

CoCo: Coherence-Enhanced Machine-Generated Text Detection Under Low Resource With Contrastive Learning

Anonymous ACL submission

Abstract

Machine-Generated Text (MGT) detection, a task that discriminates MGT from Human-Written Text (HWT), plays a crucial role in preventing misuse of text generative models, which excel in mimicking human writing style recently. Latest proposed detectors usually take coarse text sequences as input and fine-tune pretrained models with standard cross-entropy loss. However, these methods fail to consider the linguistic structure of texts. Moreover, they lack the ability to handle the low-resource problem which could often happen in practice considering the enormous amount of textual data online. In this paper, we present a coherence-based contrastive learning model named CoCo to detect the possible MGT under low-resource scenario. To exploit the linguistic feature, we encode coherence information in form of graph into text representation. To tackle the challenges of low data resource, we employ a contrastive learning framework and propose an improved contrastive loss for preventing performance degradation brought by simple samples. The experiment results on two public datasets and two self-constructed datasets prove our approach outperforms the state-of-art methods significantly.

1 Introduction

Thriving progress in the field of text generative models (TGMs) (Kenton and Toutanova, 2019; Yang et al., 2019; Liu et al., 2019; See et al., 2019; Keskar et al., 2019; Dathathri et al., 2019; Lewis et al., 2020; Brown et al., 2020; Gao et al., 2021a; Madotto et al., 2021; Ouyang et al., 2022), *e.g.*, ChatGPT¹, enables everyone to produce MGTs massively and rapidly. However, the accessibility to high-quality TGMs is prone to cause misuses, such as fake news generation (Zellers et al., 2019; Brown et al., 2020; Yanagi et al., 2020), product review forging (Adelani et al., 2020), and spamming

¹<https://chat.openai.com>

How to find hidden cameras in your Airbnb, and anywhere else		
	Human-written text	Machine-generated text
Document	S1: In recent months there's been a number of alarming reports of Airbnb hosts installing hidden cameras in their properties but not disclosing them to the guests staying there. S2: Back in January Fast Company reported on a computer science professor at Carnegie Mellon University who discovered two hidden cameras recording him and his family in an Airbnb. S3: And just last month The Atlantic reported on a New Zealand family who was renting an Airbnb in Ireland and found they were being live-streamed from a hidden security camera.	S1: Anyone who finds a video of someone on Airbnb will probably fall under the new category of hidden cameras, which can be found only in a large part of every Airbnb listing, and you're never alone. S2: Apple, Google, and Amazon combined to find the most hidden camera listings in December 2018. S3: The electronics giant's Facebook, the mapping app and the mobile messaging company Linea formed an OfficeTeam unit that can find the video even if someone's not using them, and can track real-time activity.
Sentence Interaction		

Figure 1: Illustration of sentence-level structure difference between HWT and MGT, the MGT is generated by GROVER (Zellers et al., 2019). HWT is more coherent than MGT as the sentences share more same entities with each other.

(Tan et al., 2012), etc. MGTs are hard to distinguish by an untrained human for their human-like writing style (Ippolito et al., 2020) and the excessive amount (Grinberg et al., 2019), which calls for the study of reliable automatic MGT detectors.

Previous works on MGTs detection mainly concentrate on sequence feature representation and classification (Gehrmann et al., 2019; Solaiman et al., 2019; Zellers et al., 2019). Recent studies have shown the good performance of automated detectors in a fine-tuning fashion (Solaiman et al., 2019). Although the fine-tuning based detectors have demonstrated their effectiveness, they still suffer from two issues that limit their conversion to practical use: (1) Existing detectors treat input documents as flat sequences of tokens and use neural encoders or statistical features (*e.g.*, TF-IDF) to represent text as the dense vector for classification. These methods rely much on the token-level dis-

tribution difference of texts in each class, which ignores high-level linguistic representation of text structure. (2) Compared with the enormous number of online texts, annotated dataset for training MGT detectors is rather low-resource. Constrained by the amount of available annotated data, traditional detectors sustain frustrating accuracy and even collapse during the test stage.

As shown in Fig. 1, MGTs and HWTs exhibit difference in terms of coherence traced by entity consistency. Thus, we propose an entity coherence graph to model the sentence-level structure of texts based on the thoughts of Centering Theory (Grosz and Sidner, 1986), which evaluates text coherence by entity consistency. Entity coherence graph treats entities as nodes and builds edges between entities in the same sentences and same entities among different sentences to reveal the text structure. Instead of treating text as flat sequence, coherence modeling helps to introduce distinguishable linguistic feature at input stage and provides explainable difference between MGTs and HWTs.

To alleviate the low-resource problem in the second issue, inspired by the resurgence of contrastive learning (He et al., 2020; Chen et al., 2020), we utilize proper design of sample pair and contrastive process to learn fine-grained instance-level features under low resource. However, it has been proven that the easiest negative samples are unnecessary and insufficient for model training in contrastive learning (Cai et al., 2020). To circumvent the performance degradation brought by the easy samples, we propose a novel contrastive loss with capability to reweight the effect of negative samples by difficulty score to help model concentrate more on hard samples and ignore the easy samples.

Extensive experiments on multiple datasets (GROVER, GPT-2, GPT-3) and a case study with ChatGPT-generated texts demonstrate the effectiveness of our proposed method.

In summary, our contributions are summarized as follows:

- **Coherence Graph Construction:** We model the text coherence with entity consistency and sentence interaction while statistically proving its distinctiveness in MGTs detection, and further introduce this linguistic feature at input stage.
- **Improved Contrastive Loss:** We propose a novel contrastive loss in which hard negative

samples are paid more attention for improving detection accuracy of challenging sample.

- **Outstanding Performance:** We achieve state-of-art performance on four MGT datasets in both low-resource and high-resource setting. Experimental results verify the effectiveness of our model.

2 Related Work

Machine-generated Text Detection. Machine-generated texts, also named deepfake or neural fake texts, are generated by language models to mimic human writing style, making them perplexing for humans to distinguish (Ippolito et al., 2020). Generative models like GROVER (Zellers et al., 2019), GPT-2 (Radford et al., 2019), GPT-3 (Brown et al., 2020) and emerging ChatGPT has been evaluated on the MGT detection task and achieve good results. Bakhtin et al. (2019) train an energy-based model by treating the output of TGMs as negative samples to demonstrate the generalization ability. Deep learning models that incorporate stylometry and external knowledge are also feasible for improving the performance of MGT detectors (Uchendu et al., 2019; Zhong et al., 2020). Our method differs from the previous work by analyzing and modeling text coherence as distinguishable feature and emphasizing the performance improvement under low-resource scenario.

Coherence Modeling. For generative models, coherence is the critical requirement and vital target (Hovy, 1988). Previous works mainly discuss two types of coherence, local coherence (Mellish et al., 1998; Althaus et al., 2004) and global coherence (Mann and Thompson, 1987). Local coherence focus on sentence-to-sentence transitions (Lapata, 2003), while global coherence tries to capture comprehensive structure (Karamanis and Manurung, 2002). Our method strives to represent both local and global coherence with inner- and inter-sentence relations between entity nodes.

Contrastive Learning. Contrastive learning in NLP demonstrates superb performance in learning token-level embeddings (Su et al., 2021) and sentence-level embeddings (Gao et al., 2021b) for natural language understanding. With in-depth study of the mechanism of contrastive learning, the hardness of samples is proved to be crucial in the training stage. Cai et al. (2020) define the dot product between the queries and the negatives

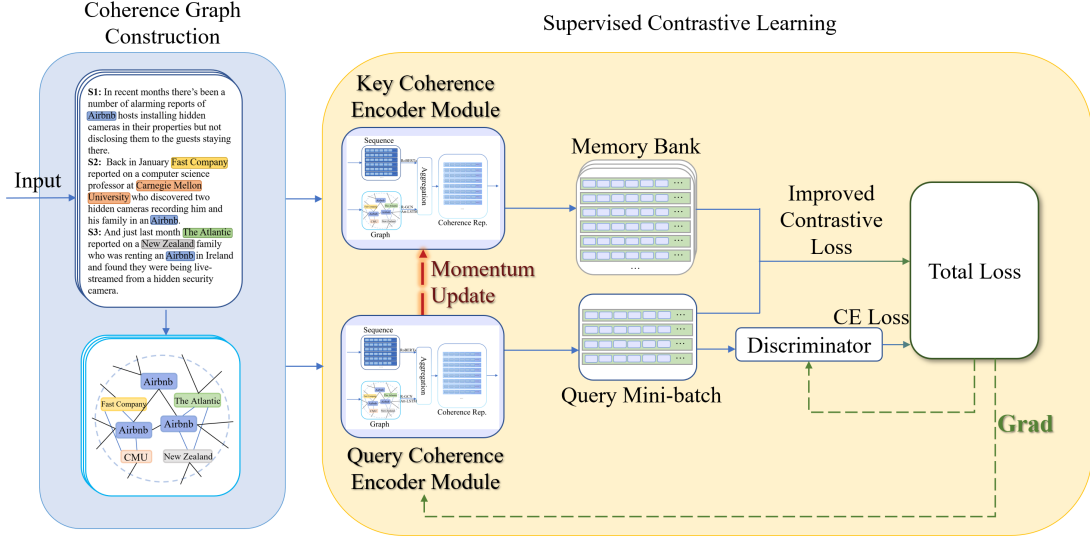


Figure 2: Overview of CoCo. Input document is parsed to construct a coherence graph (3.1), the text and graph are utilized by a supervised contrastive learning framework (3.2), in which coherence encoding module is designed to encode and aggregate to generate coherence-enhanced representation (3.2.3). After that, we employ a MoCo based contrastive learning architecture in which key encodings are stored in a dynamic memory bank (3.2.4) with improved contrastive loss to make final prediction (3.2.5).

in normalized embedding space as hardness and figured out the easiest 95% negatives are insufficient and unnecessary. Song et al. (2022) propose a difficulty measure function based on the distance between classes and apply curriculum learning to the sampling stage. Differently, our method pays more attention to hard negative samples for improving the detection accuracy of challenging samples.

3 Methodology

The workflow of CoCo mainly contains coherence graph construction, and supervised contrastive learning discriminator, and Fig. 2 illustrates its overall architecture.

3.1 Coherence Graph Construction

In this part, we illustrate how to construct coherence graph to dig out coherence structure of text by modeling sentence interaction.

According to Centering Theory (Grosz and Sidner, 1986), coherence of texts could be modeled by sentence interaction around center entities. To better reflect text structure and avoid semantic overlap, we proposed to construct an undirected graph with entities as nodes. Specifically, we first implement the ELMo-based NER model TagLM (Peters et al., 2017) with the help of NER toolkit AllenNLP²

²<https://demo.allennlp.org/named-entity-recognition>

to extract the entities from document. An relation $\langle inter \rangle$ is constructed between same entities in different sentences and nodes within same sentences are connected by relation $\langle inner \rangle$ for their natural structure relevance. Formally, the mathematical form of coherence graph’s adjacent matrix is defined as follows:

$$A_{ij} = \begin{cases} 1 & \text{rel } \langle inner \rangle & v_{i,a} \neq v_{j,b}, a = b \\ 1 & \text{rel } \langle inter \rangle & v_{i,a} = v_{j,b}, a \neq b \\ 0 & \text{rel } \text{None} & \text{others} \end{cases}$$

where $v_{i,a}$ represents i -th entity in sentence a , which is regarded as node in coherence graph.

3.2 Supervised Contrastive Learning

3.2.1 Model Overview

The training process is illustrated in Fig. 2. Each entry in the dataset is document with its coherence graph. The entries in training set are sampled randomly into keys and queries. Two coherence encoder modules (CEM) f_k and f_q , are initialized the same to generate coherence-enhanced representation D_k and D_q for key and query. A dynamic memory bank with the size of all training data is initialized to store all key representation and their annotations for providing enough contrastive pairs in low-resource scenario. In every training step, the newly encoded key graphs update memory bank following First In First Out (FIFO) rule to keep it updated and the training process consistent. A

novel loss composed of improved contrastive loss and cross-entropy loss ensures the model’s ability to achieve instance-level intra-class compactness and inter-class separability while maintaining the class-level distinguishability. A linear discriminator takes query representations as input and generates prediction results. The pseudocode of training process is shown in Appendix A.8.

3.2.2 Positive/Negative Pair Definition

In supervised setting, where we have access to label information, we define two samples with same label as positive pair and that with different labels as negative pair for incorporating label information into training process.

3.2.3 Encoder Design

In this part, we introduce how to initialize node representation and graph neural network structure which is utilized to integrate coherence information into semantic representation of text by propagating and aggregating information from different granularity with an innovated coherence encoder module (CEM).

Node Representation Initialization. We initialize the representation of entity nodes with powerful pre-trained model RoBERTa for its superior ability to encode contextual information into text representation.

Given an entity e with a span of n tokens, we utilize RoBERTa to map input document x to embeddings $h(x)$. The contextual representation of e is calculated as follows:

$$Z_v = \frac{1}{n} \sum_{i=0}^n h(x)_{e_i}, \quad (1)$$

where e_i is the absolute position where the i -th token in e lies in the whole document.

Relation-aware GCN. Based on the vanilla Graph Convolutional Networks (Welling and Kipf, 2016), we propose a novel method to assign different weight W_r for inter and inner relation r with Relation-aware GCN. Relation-aware GCN convolute edges of each kind of relation in the coherence graph separately. The final representation is the sum of GCN outputs from all relations. We use two-layer GCN in the model because more layers will cause an overfitting problem under low resources. We define the relation set as R , and the calculation formula is as follows:

$$H^{(i+1)} = \sum_{r \in R} \hat{A} \text{ReLU}((\hat{A} H^{(i)} W_r^{(i)}) W_r^{(i+1)}), \quad (2)$$

$$\hat{A} = \tilde{D}^{-\frac{1}{2}} \tilde{A} \tilde{D}^{-\frac{1}{2}},$$

where $H^{(i)} \in \mathbb{R}^{N \times d}$ is node representation in i -th layer. $\tilde{A} = A + I$, A is the adjacency matrix of the coherence graph, $\hat{A}$ is the normalized Laplacian matrix of $\tilde{A}$, W_r is the relation transformation matrix for relation r .

Sentence Representation. Afterward, we aggregate updated node representation from last layer of Relation-aware GCN into sentence-level representation to prepare for concatenation with sequence representation from RoBERTa. The aggregation follows the below rule:

$$Z_{s_i} = \frac{1}{M_i} \sum_j^{M_i} \sigma(W_s H_{(i,j)} + b_s), \quad (3)$$

where M_i represents the number of entities in i -th sentence, $H_{(i,j)}$ represents the embedding of j -th entity in i -th sentence, W_s is weight matrix and b_s is bias. All the sentence representations within same document are concatenated as sentence matrix Z_s .

Document Representation with Attention LSTM. We design a self-attention mechanism for discovering the sentence-level coherence between one sentence and other sentences, and apply LSTM with the objective to track the coherence in continuous sentences and take the last hidden state of LSTM for aggregated document representation containing comprehensive coherence information. The calculation is described as follows:

$$Z_c = \text{LSTM}(\text{softmax}(\gamma \frac{\text{norm}(K) \text{norm}(Q)^T}{\sqrt{d_Z}}) V), \quad (4)$$

where K, Q, V are linear transformations of Z_s with matrix W_k, W_q, W_v , d_Z is the dimension of representation Z_s , and γ is a hypergamma-parameter for scaling.

Finally, we concatenate Z_c and the sequence representation $h([\text{CLS}])$ from the RoBERTa’s last layer to generate document coherence-enhanced representation D .

3.2.4 Dynamic Memory Bank

The dynamic memory bank is created to store as much as key encoding D_k to form adequate positive and negative pairs within a batch. The dynamic

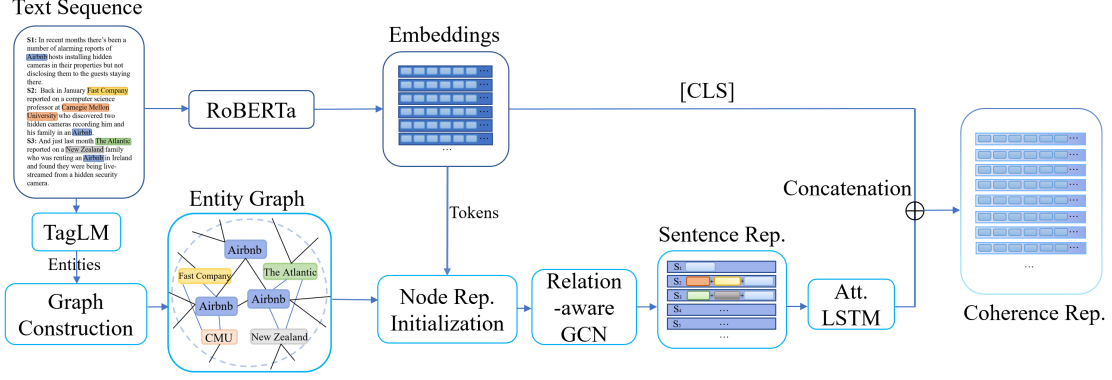


Figure 3: Illustration of coherence encoder module (CEM) which encodes and fuses the coherence graph and text sequence to generate coherence-enhanced representation of document.

memory bank is maintained as a queue so that the newly encoded keys could replace the outdated ones, which keeps the consistency between the key encoding and current training step.

3.2.5 Loss Function

Following the definition of positive pairs and negative pairs above, traditional supervised contrastive loss (Gunel et al., 2020) treats all positive pairs and negative pairs equally.

However, with recognition that not all negatives are created equal (Cai et al., 2020), our goal is to emphasize the informative samples for helping the model to differentiate difficult samples. Thus, we propose an improved contrastive loss which dynamically adjusts the weight of negative pair similarity according to the hardness of negative samples. To be specific, the hard negative samples should be assigned larger weight for stimulating the model to better pull same class together and push different class away. The improved contrastive loss is defined as:

$$\mathcal{L}_{ICL} = \sum_{j=1}^M \mathbf{1}_{y_i=y_j} \log \frac{S_{ij}}{\sum_{p \in \mathcal{P}(i)} S_{ip} + \sum_{n \in \mathcal{N}(i)} r f_{in} S_{in}},$$

$$r f_{ij} = \beta \frac{D_q^i D_k^n}{\text{avg}(D_q^i D_k^{1:|\mathcal{N}(i)|})},$$

$$S_{ij} = \exp(D_q^i D_k^j / \tau),$$
(5)

where $\mathcal{P}(i)$ is the positive set in which data has the same label with q_i and $\mathcal{N}(i)$ is the negative set in which data has different label from q_i .

Apart from instance-level learning mechanism, a linear classifier combined with cross entropy loss $\mathcal{L}_{CE}$ is employed to provide the model with class-

level separation ability. $\mathcal{L}_{CE}$ is calculated by

$$\mathcal{L}_{CE} = \frac{1}{N} \sum_{i=1}^N -[y_i \log(p_i) + (1-y_i) \log(1-p_i)],$$
(6)

where p_i is the prediction probability distribution of i -th sample. The final loss $\mathcal{L}_{total}$ is a weighted average of $\mathcal{L}_{ICL}$ and $\mathcal{L}_{CE}$ as:

$$\mathcal{L}_{total} = \alpha \mathcal{L}_{ICL} + (1 - \alpha) \mathcal{L}_{CE},$$
(7)

where the hyperparameter α adjusts the relative balance between instance compactness and class separability.

3.2.6 Momentum Update

The parameters of query encoder f_q and the classifier can be updated by gradient back-propagated from $\mathcal{L}_{total}$. We denote the parameters of f_q as θ_q , the parameters of f_k as θ_k . The key encoder f_k 's parameters are updated by momentum update mechanism:

$$\theta_k \leftarrow \beta \theta_k + (1 - \beta) \theta_q,$$
(8)

where the hyperparameter β is momentum coefficient.

4 Experiments

4.1 Datasets

We evaluate our model on the following datasets:

- **GROVER Dataset** is a News-Style dataset provided by (Zellers et al., 2019), in which HWTs are collected from RealNews, a large corpus of news articles from Common Crawl, and MGTs are generated by Grover-Mega, a transformer-based news generator with parameters sized 1.5B.

Dataset		Train	Valid	Test
GROVER	HWT	5,000	2,000	8,000
	MGT	5,000	1,000	4,000
GPT-2	HWT	25,000	5,000	5,000
	MGT	25,000	5,000	5,000
GPT-3 unmixed	HWT	5,455	1,000	1,000
	MGT	5,455	1,000	1,000
GPT-3 mixed	HWT	5,033	1,000	1,000
	MGT	5,033	1,000	1,000

Table 1: Basic statistics of datasets.

- **GPT-2 Dataset** is a Webtext-style dataset provided by OpenAI³ with HWTs adopted from WebText and MGTs produced by GPT-2 XLM-1542M.
- **GPT-3 Dataset** is a News-Style open source dataset constructed by us based on the text-davinci-003⁴ model of OpenAI, which is the most capable GPT-3 model so far and can generate longer texts (maximum 4,000 tokens). The GPT-3 model refer to various latest newspapers (Dec. 2022 - Present) whose full texts act as the HWTs part and generate news by imitation. We use two subsets: mixed- and unmixed-provenances. The details of this dataset are explained in Appendix A.1.

The statistics of datasets is summarized in Table 1 and the implementation details are in Appendix A.3.

4.2 Comparison Models

We compare CoCo to the state-of-art detection methods to reveal the effectiveness. The baselines are introduced as follows.

- **GPT-2** (Radford et al., 2019), a strong pre-trained language model based on the decoder architecture of transformer. We use GPT-2 small with parameters sized 124M for the fairness of comparison.
- **RoBERTa** (Liu et al., 2019), a powerful transformers-based bidirectional language model with robust performance on downstream tasks. We use RoBERTa-base sized 110M in the experiment.
- **XLNet** (Yang et al., 2019), a language model which is superb in understanding long docu-

ments. We exploit XLNet-base whose scale is 110M.

- **CE+SCL** (Gunel et al., 2020), a state-of-the-art supervised contrastive learning method in various downstream task. We train the detector with Cross-Entropy loss (CE) and supervised contrastive loss (SCL) calculated within a mini-batch.
- **DualCL** (Chen et al., 2022), a contrastive learning method with the addition of label representations for data augmentation.

4.3 Performance Comparison

As shown in Table 2, CoCo surpasses state-of-the-art methods in MGT detection task 4.25%, 5.98% (limited dataset), and 3.31%, 3.58% (full dataset) in average in terms of accuracy and F1-score, respectively. The result indicates the rationality of coherence graph designing and the effectiveness of encoding coherence information into document representation, which further proves the coherence difference between MGT and HWT. Moreover, it should be noticed that CoCo gets less affected by randomness, which illustrates that the coherence graph we construct is a robust feature that helps stabilize the model performance. Meanwhile, CoCo outperforms CE+SCL and DualCL regardless the size of training set, which suggests the success of improved contrastive loss to solve the performance degradation problem brought by simple negative samples.

We also find GROVER Dataset is the hardest to detect. It is because GROVER generator is trained in an adversarial fashion with the objective to deceive the verifier, which endows the generator with deceptive nature. To our surprise, GPT-3 dataset is overly simple to all detectors. We conduct extensive experiments on different self-constructed and published GPT-3 dataset generated by a series prompts, which also validate this thundering conclusion. The experiment details and results are in Appendix A.2.

Due to the secrecy of GPT-3 implementation details and the lack of interpretability of LLMs, we can only make several inferences for this phenomenon up to the best of our knowledge. First, since GPT-3 is fine-tuned with RLHF (Ouyang et al., 2022), we suppose that instructions like "intimate a piece of news" is rare in prompt dataset used in GPT-3 fine-tuning, which causes its inadequate ability to imitate news. In fact, we find news

³<https://github.com/openai/gpt-2-output-dataset>

⁴<https://beta.openai.com/docs/models/gpt-3>

Dataset Size Metric	GROVER				GPT-2			
	Limited Dataset (10%)		Full Dataset		Limited Dataset (10%)		Full Dataset	
	ACC	F1	ACC	F1	ACC	F1	ACC	F1
GPT2	0.6401 $\pm$ 0.0136	0.4289 $\pm$ 0.0413	0.8274 $\pm$ 0.0091	0.8003 $\pm$ 0.0141	0.8575 $\pm$ 0.0041	0.8406 $\pm$ 0.0070	0.8913 $\pm$ 0.0066	0.8839 $\pm$ 0.0078
XLNet	0.6906 $\pm$ 0.0321	0.5193 $\pm$ 0.0365	0.8156 $\pm$ 0.0079	0.7493 $\pm$ 0.0073	0.8837 $\pm$ 0.0031	0.8732 $\pm$ 0.0041	0.9091 $\pm$ 0.0091	0.9027 $\pm$ 0.0111
RoBERTa	0.7730 $\pm$ 0.0121	0.6370 $\pm$ 0.0186	0.8772 $\pm$ 0.0029	0.8171 $\pm$ 0.0048	0.9244 $\pm$ 0.0041	0.9214 $\pm$ 0.0045	0.9402 $\pm$ 0.0039	0.9384 $\pm$ 0.0044
DualCL	0.6926 $\pm$ 0.0634	0.5397 $\pm$ 0.0971	0.7574 $\pm$ 0.0855	0.6388 $\pm$ 0.130	0.7874 $\pm$ 0.1210	0.6950 $\pm$ 0.0944	0.8023 $\pm$ 0.1120	0.8046 $\pm$ 0.1530
CE+SCL	0.7777 $\pm$ 0.0134	0.6447 $\pm$ 0.0286	0.8782 $\pm$ 0.0044	0.8202 $\pm$ 0.0057	0.9255 $\pm$ 0.0030	0.9241 $\pm$ 0.0016	0.9408 $\pm$ 0.0006	0.9390 $\pm$ 0.0009
CoCo	0.7808 $\pm$ 0.0044	0.6543 $\pm$ 0.0106	0.8826 $\pm$ 0.0018	0.8265 $\pm$ 0.0036	0.9271 $\pm$ 0.0019	0.9254 $\pm$ 0.0018	0.9457 $\pm$ 0.0004	0.9452 $\pm$ 0.0004

Dataset Size Metric	GPT-3 Unmixed				GPT-3 Mixed			
	Limited Dataset (10%)		Full Dataset		Limited Dataset (10%)		Full Dataset	
	ACC	F1	ACC	F1	ACC	F1	ACC	F1
GPT2	0.9631 $\pm$ 0.0037	0.9666 $\pm$ 0.0044	0.9917 $\pm$ 0.0056	0.9905 $\pm$ 0.0042	0.9623 $\pm$ 0.0039	0.9636 $\pm$ 0.0046	0.9910 $\pm$ 0.0046	0.9910 $\pm$ 0.0033
XLNet	0.9252 $\pm$ 0.0052	0.9241 $\pm$ 0.0054	0.9620 $\pm$ 0.0043	0.9634 $\pm$ 0.0068	0.9149 $\pm$ 0.0044	0.9152 $\pm$ 0.0059	0.9513 $\pm$ 0.0052	0.9505 $\pm$ 0.0039
RoBERTa	0.9916 $\pm$ 0.0043	0.9899 $\pm$ 0.0049	0.9928 $\pm$ 0.0035	0.9913 $\pm$ 0.0040	0.9892 $\pm$ 0.0022	0.9893 $\pm$ 0.0025	0.9923 $\pm$ 0.0017	0.9901 $\pm$ 0.0024
CE+SCL	0.9918 $\pm$ 0.0039	0.9901 $\pm$ 0.0058	0.9944 $\pm$ 0.0023	0.9943 $\pm$ 0.0031	0.9903 $\pm$ 0.0020	0.9898 $\pm$ 0.0023	0.9932 $\pm$ 0.0017	0.9905 $\pm$ 0.0038
CoCo	0.9932 $\pm$ 0.0042	0.9913 $\pm$ 0.0033	0.9972 $\pm$ 0.0015	0.9957 $\pm$ 0.0020	0.9913 $\pm$ 0.0034	0.9911 $\pm$ 0.0038	0.9932 $\pm$ 0.0019	0.9937 $\pm$ 0.0028

Table 2: Results of the model comparison. It should be noticed that DualCL is easily affected by random seed, which may be caused by its weakness in understanding long texts. We do not present the experiment results for DualCL on GPT-3 dataset because the documents in GPT-3 dataset is so long that DualCL completely fails.

generated by GPT-3 is more like a short version of reference news (far shorter than the max output length) instead of the imitation. Second, the HWTs we use are latest news which are not included in GPT-3’s training data. The lack of knowledge background about reference news might limit GPT-3’s associative ability. Third, we do not fine-tune GPT-3 with HWTs we collect while model used in GPT-2 dataset did. The unfamiliarity of GPT-3 with the texts it is required to imitate could impair its ability to rewrite news in a similar style. We will further investigate our hypotheses and provide possible deep insights into this counterintuitive but very interesting phenomenon, which may also exist in the GPT-4 model, in future works.

4.4 Ablation Study

To illustrate the necessity of some components of CoCo, we conduct several ablation experiments on limited GROVER dataset and the results are shown in Table 3. We introduce the ablation model structure below:

Model	ACC	F1
CoCo (Plain)	0.7697	0.6428
CoCo (Sentence nodes)	0.7733	0.6379
CoCo (Coherence)	0.7777	0.6463
CoCo (Coherence + LSTM)	0.7787	0.6471
CoCo (Coherence + LSTM + SCL)	0.7827	0.6609
CoCo	0.7843	0.6684

Table 3: Results of ablation study.

CoCo (Plain) removes graph information and encodes only by RoBERTa parts. Moreover, the model removes contrastive learning and uses CE loss.

CoCo (Sentence Nodes) treats sentences as nodes instead of entities and establish edges between sentences which share same entities. Node representation is initialized by RoBERTa embedding and mean-pooling operation. Document representation is obtained by one CEM discarding sentence representation and attention LSTM part in section 3.2.3. Document representation is calculated by mean-pooling operation on sentence node representations. A linear classification head with cross-entropy loss is used for detection.

CoCo (Coherence) incorporates coherence graph into the representation of document and deploys sentence representation part in section 3.2.3. The rest are the same with CoCo (Sentence Nodes).

CoCo (Coherence + LSTM) uses attention LSTM for document-level aggregation and the rest is the same as CoCo (Coherence).

CoCo (Coherence + LSTM + SCL) utilizes the contrastive learning framework but the loss function is traditional supervised contrastive loss instead of the improved contrastive loss.

As shown in Table 3, coherence information and the contrastive learning framework greatly contribute to the development of model performance, especially in F1-Score. Replacing entity nodes in

coherence graph with sentences impairs the detector, which could be caused by semantic overlap between graph representation and text sequence representation. The attention LSTM also plays an important role in preserving coherence information during sentence aggregation. Lastly, the results also shows the advantage of improved contrastive loss over standard supervised contrastive loss. We also investigate the effect of hyperparameters on CoCo, the results are shown in Appendix A.4.

4.5 Discussion: Utility of Coherence Graph

4.5.1 Case Study

In this subsection, we conduct a case study with HWT and MGT produced by sensational ChatGPT with the same metadata. As illustrated in Fig. 4, we parse two news as coherence graphs. And we observe that although ChatGPT expresses fluently, it is not coherent from the perspective of coherence graph. Hence, CoCo utilizes the distinctive coherence feature and makes correct predictions while RoBERTa fails. This reflects even the most popular and advanced language model could suffer from weak coherence and be detected by CoCo.

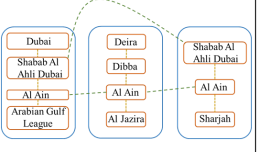
Shabab Al Ahli Dubai have task cut out		
	HWT: Gulf News	MGT: ChatGPT
Document	<i>Dubai: Shabab Al Ahli Dubai will continue pushing towards the top while keeping an eye on chasing Al Ain as the 20th round of the Arabian Gulf League (AGL) is played out this weekend. The team from Deira climbed into second place for the first time last week following their 3-1 win over Dibba and a 5-1 thrashing of Al Ain by Al Jazira. Shabab Al Ahli Dubai now remain tied on 38 points with defending champions Al Ain, while unbeaten Sharjah are at the top with 47 points.</i>	<i>Shabab Al Ahli Dubai FC, the Dubai-based football club, have a difficult task ahead of them as they gear up for the upcoming season. The team, which plays in the UAE Pro League, is facing a number of challenges, both on and off the field, that will need to be overcome if they hope to be successful in the coming months. In the off-season, several of the team's top performers, including star striker Ahmed Khalil and midfield maestro Omar Abdulrahman, left the club to join teams in other countries.</i>
Coherence Graph		
	CoCo: [0.9233, 0.0767] RoBERTa: [0.9085, 0.0915]	CoCo: [0.1038, 0.8962] RoBERTa: [0.7354, 0.2646]

Figure 4: An illustration for case study of our method. Entities in documents are colored green. The blue solid box indicates the sentence. The orange dashed lines are inner edges and green dashed lines are inter edges. Numbers in red indicate the probability of predicted label.

4.5.2 Static Analysis for Coherence Graph

We apply static geometric features analysis on coherence graph we construct to provide an in-depth view on linguistic explanation about graph structure. In the following discussion, we take the

dataset of GROVER into the analysis. Some basic metrics of data and the corresponding graph are shown in Table 9.

Metric	HWT	MGT
Sample Num.	4994	4991
Avg. Num. of Token	463.2	456.0
Avg. Num. of Vertex	43.60	32.37
Avg. Num. of Edge	107.4	65.44

Table 4: Basic metrics of texts and corresponding coherence graphs.

Degree Distribution of coherence graph semantically measures the co-occurrence and TF-IDF feature of keywords, showing global coherence because high-degree nodes devote to the main topic and low-degree nodes are the extension. The degree of the graph representation of HWTs is **2.980**, which is **15.0%** larger than MGTs (2.591), and the distribution of HWTs has a longer tail than MGTs. Furthermore, we prove that degree distribution can robustly detect MGTs and HWTs when impacted by style and genre differences. More details are discussed in the Appendix A.5.

5 Conclusion

In this paper, we propose CoCo, a coherence-enhanced contrastive learning model for MGT detection. We construct a novel coherence graph from document and implement a MoCo-based contrastive learning framework to improve model performance in low-resource setting. An innovative encoder composed of relation-aware GCN and attention LSTM is designed to learn the coherence representation from coherence graph which is further incorporated with sequence representation of document. To alleviate the effect of unnecessary easy samples, we propose an improved contrastive learning loss to force the model to pay more attention to hard negative samples. We evaluate our method on MGT datasets generated by GROVER, GPT-2, and GPT-3, respectively, in both low-resource and high-resource settings. CoCo outperforms Transformer-based methods and contrastive-learning-based methods on all datasets and both settings.

Acknowledgements

We thank the six anonymous reviewers and the area chair for their helpful comments and feedback, which aided us in greatly improving the paper.

Limitations

In this work, we step forward to better distinguishing MGTs under the low-resource setting. However, several limitations still exist for the broader applications of this detector. Firstly, MGTs are easier to generate and collect than HWTs, which may cause an imbalanced label distribution in the dataset. And CoCo literally corrupts in extremely imbalanced data distribution condition, as shown in A.6. Future work could build upon the contrastive learning method of CoCo with innovation on sampling strategy for harsh low-resource and imbalanced data settings. Secondly, our method artificially generates a coherence graph for every entry, which is not efficient for larger datasets. What's more, short text, codes, and mathematical proofs, which are hard to generate coherence graphs, are also limitedly detected by CoCo. More distinctive and easy-to-calculate features are worth exploring for generating distinguishable representations for texts with efficiency while better understanding the essence of TGMs. Thirdly, with instruct-based generation and human-in-loop fine-tuning models prevailing, the strategy and defect of TGMs change slightly but constantly. The entity relation with the same semantic granularity and concretization in this paper would not be enough to detect the high-quality content by TGMs in the future. More generative and adaptive detection models should be considered.

Ethical Considerations

We provide insight into the potential weakness of TGMs and publish GPT-3 news dataset. We understand that the discovery of our work can be viciously used to confront detectors. And we understand that malicious users can copy the contents of our GPT-3 news dataset to disguise real news and publish them. However, with the purpose of calling for attention to detecting and controlling possible misuse of TGMs, we believe our work will inspire the advance of the stronger detector of MGTs and prevent all potential negative uses of language models.

Our work complies with sharing & publication policy of OpenAI⁵ and all data we collect is in public domain and licensed for research purposes.

⁵<https://openai.com/api/policies/sharing-publication/>

References

- David Ifeoluwa Adelani, Haotian Mai, Fuming Fang, Huy H Nguyen, Junichi Yamagishi, and Isao Echizen. 2020. Generating sentiment-preserving fake online reviews using neural language models and their human-and machine-based detection. In *International Conference on Advanced Information Networking and Applications*, pages 1341–1354. Springer.
- Ernst Althaus, Nikiforos Karamanis, and Alexander Koller. 2004. Computing locally coherent discourses. In *Proceedings of the 42nd Annual Meeting of the Association for Computational Linguistics (ACL-04)*, pages 399–406.
- Anton Bakhtin, Sam Gross, Myle Ott, Yuntian Deng, Marc'Aurelio Ranzato, and Arthur Szlam. 2019. Real or fake? learning to discriminate machine from human generated text. *arXiv preprint arXiv:1906.03351*.
- Roi Blanco and Christina Lioma. 2011. Graph-based term weighting for information retrieval. *Information Retrieval*, 15:54–92.
- Stefan Bordag, Gerhard Heyer, and Uwe Quasthoff. 2003. Small worlds of concepts and other principles of semantic search. In *International Workshop on Innovative Internet Community Systems*, pages 10–19. Springer.
- Tom Brown, Benjamin Mann, Nick Ryder, Melanie Subbiah, Jared D Kaplan, Prafulla Dhariwal, Arvind Neelakantan, Pranav Shyam, Girish Sastry, Amanda Askell, et al. 2020. Language models are few-shot learners. *Advances in neural information processing systems*, 33:1877–1901.
- Tiffany Tianhui Cai, Jonathan Frankle, David J Schwab, and Ari S Morcos. 2020. Are all negatives created equal in contrastive instance discrimination? *arXiv preprint arXiv:2010.06682*.
- Ramon Ferrer Cancho and Ricard V. Solé. 2001. Two regimes in the frequency of words and the origins of complex lexicons: Zipf's law revisited*. *Journal of Quantitative Linguistics*, 8:165 – 173.
- Ramon Ferrer Cancho, Ricard V Solé, and Reinhard Köhler. 2004. Patterns in syntactic dependency networks. *Physical Review E*, 69(5):051915.
- Qianben Chen, Richong Zhang, Yaowei Zheng, and Yongyi Mao. 2022. Dual contrastive learning: Text classification via label-aware data augmentation. *arXiv preprint arXiv:2201.08702*.
- Xinlei Chen, Haoqi Fan, Ross Girshick, and Kaiming He. 2020. Improved baselines with momentum contrastive learning. *arXiv preprint arXiv:2003.04297*.
- Monojit Choudhury, Markose Thomas, Animesh Mukherjee, Anupam Basu, and Niloy Ganguly. 2007. How difficult is it to develop a perfect spell-checker? a cross-linguistic analysis through complex network

656	approach. In <i>Proceedings of the Second Workshop on TextGraphs: Graph-Based Algorithms for Natural Language Processing</i> , pages 81–88.	
657		
658		
659	Sumanth Dathathri, Andrea Madotto, Janice Lan, Jane Hung, Eric Frank, Piero Molino, Jason Yosinski, and Rosanne Liu. 2019. Plug and play language models: A simple approach to controlled text generation. <i>arXiv preprint arXiv:1912.02164</i> .	
660		
661		
662		
663		
664	Sergey N Dorogovtsev and José Fernando F Mendes. 2001. Language as an evolving word web. <i>Proceedings of the Royal Society of London. Series B: Biological Sciences</i> , 268(1485):2603–2606.	
665		
666		
667		
668	Ramon Ferrer Cancho, Andrea Capocci, and Guido Caldarelli. 2007. Spectral methods cluster words of the same class in a syntactic dependency network. <i>International journal of bifurcation and chaos</i> , 17(7):2453–2463.	
669		
670		
671		
672		
673	Tianyu Gao, Adam Fisch, and Danqi Chen. 2021a. Making pre-trained language models better few-shot learners. In <i>Proceedings of the 59th Annual Meeting of the Association for Computational Linguistics and the 11th International Joint Conference on Natural Language Processing (Volume 1: Long Papers)</i> , pages 3816–3830.	
674		
675		
676		
677		
678		
679		
680	Tianyu Gao, Xingcheng Yao, and Danqi Chen. 2021b. Simcse: Simple contrastive learning of sentence embeddings. In <i>Proceedings of the 2021 Conference on Empirical Methods in Natural Language Processing</i> , pages 6894–6910.	
681		
682		
683		
684		
685	Sebastian Gehrmann, Hendrik Strobelt, and Alexander M Rush. 2019. Gltr: Statistical detection and visualization of generated text. In <i>Proceedings of the 57th Annual Meeting of the Association for Computational Linguistics: System Demonstrations</i> , pages 111–116.	
686		
687		
688		
689		
690		
691	Nir Grinberg, Kenneth Joseph, Lisa Friedland, Briony Swire-Thompson, and David Lazer. 2019. Fake news on twitter during the 2016 us presidential election. <i>Science</i> , 363(6425):374–378.	
692		
693		
694		
695	Barbara J Grosz and Candace L Sidner. 1986. Attention, intentions, and the structure of discourse. <i>Computational linguistics</i> , 12(3):175–204.	
696		
697		
698	Beliz Gunel, Jingfei Du, Alexis Conneau, and Ves Stoyanov. 2020. Supervised contrastive learning for pre-trained language model fine-tuning. <i>arXiv preprint arXiv:2011.01403</i> .	
699		
700		
701		
702	Geehan Sabah Hassan, Asma Khazaal Abdulsahib, and Siti Sakira Kamaruddin. 2017. Graph-based text representation: a survey of current approaches. <i>Research Journal of Applied Sciences, Engineering and Technology</i> , 14(9):334–340.	
703		
704		
705		
706		
707	Kaiming He, Haoqi Fan, Yuxin Wu, Saining Xie, and Ross Girshick. 2020. Momentum contrast for unsupervised visual representation learning. In <i>Proceedings of the IEEE/CVF conference on computer vision and pattern recognition</i> , pages 9729–9738.	
708		
709		
710		
711		
	Xiaochen Hou, Peng Qi, Guangtao Wang, Rex Ying, Jing Huang, Xiaodong He, and Bowen Zhou. 2021. Graph ensemble learning over multiple dependency trees for aspect-level sentiment classification. In <i>Proceedings of the 2021 Conference of the North American Chapter of the Association for Computational Linguistics: Human Language Technologies</i> , pages 2884–2894.	712 713 714 715 716 717 718 719
	Eduard H Hovy. 1988. Planning coherent multisentential text. In <i>Proceedings of the 26th annual meeting on Association for Computational Linguistics</i> , pages 163–169.	720 721 722 723
	Binxuan Huang and Kathleen M Carley. 2019. Syntax-aware aspect level sentiment classification with graph attention networks. In <i>Proceedings of the 2019 Conference on Empirical Methods in Natural Language Processing and the 9th International Joint Conference on Natural Language Processing (EMNLP-IJCNLP)</i> , pages 5469–5477.	724 725 726 727 728 729 730
	Daphne Ippolito, Daniel Duckworth, Chris Callison-Burch, and Douglas Eck. 2020. Automatic detection of generated text is easiest when humans are fooled. In <i>Proceedings of the 58th Annual Meeting of the Association for Computational Linguistics</i> , pages 1808–1822.	731 732 733 734 735 736
	Nikiforos Karamanis and Hisar Maruli Manurung. 2002. Stochastic text structuring using the principle of continuity. In <i>Proceedings of the International Natural Language Generation Conference</i> , pages 81–88.	737 738 739 740
	Jacob Devlin Ming-Wei Chang Kenton and Lee Kristina Toutanova. 2019. Bert: Pre-training of deep bidirectional transformers for language understanding. In <i>Proceedings of NAACL-HLT</i> , pages 4171–4186.	741 742 743 744
	Nitish Shirish Keskar, Bryan McCann, Lav R Varshney, Caiming Xiong, and Richard Socher. 2019. Ctrl: A conditional transformer language model for controllable generation. <i>arXiv preprint arXiv:1909.05858</i> .	745 746 747 748
	Zornitsa Kozareva, Ellen Riloff, and Eduard Hovy. 2008. Semantic class learning from the web with hyponym pattern linkage graphs. In <i>Proceedings of ACL-08: HLT</i> , pages 1048–1056.	749 750 751 752
	Mirella Lapata. 2003. Probabilistic text structuring: Experiments with sentence ordering. In <i>ACL</i> , volume 3, pages 545–552. Citeseer.	753 754 755
	Mike Lewis, Yinhan Liu, Naman Goyal, Marjan Ghazvininejad, Abdelrahman Mohamed, Omer Levy, Veselin Stoyanov, and Luke Zettlemoyer. 2020. Bart: Denoising sequence-to-sequence pre-training for natural language generation, translation, and comprehension. In <i>Proceedings of the 58th Annual Meeting of the Association for Computational Linguistics</i> , pages 7871–7880.	756 757 758 759 760 761 762 763
	Xien Liu, Xinxin You, Xiao Zhang, Ji Wu, and Ping Lv. 2020. Tensor graph convolutional networks for text classification. In <i>Proceedings of the AAAI conference on artificial intelligence</i> , volume 34, pages 8409–8416.	764 765 766 767 768

769	Yinhan Liu, Myle Ott, Naman Goyal, Jingfei Du, Man-	Abigail See, Aneesh Pappu, Rohun Saxena, Akhila	821
770	dar Joshi, Danqi Chen, Omer Levy, Mike Lewis,	Yerukola, and Christopher D Manning. 2019. Do	822
771	Luke Zettlemoyer, and Veselin Stoyanov. 2019.	massively pretrained language models make better	823
772	Roberta: A robustly optimized bert pretraining ap-	storytellers? In <i>Proceedings of the 23rd Confer-</i>	824
773	proach. <i>arXiv preprint arXiv:1907.11692</i> .	<i>ence on Computational Natural Language Learning</i>	825
		(CoNLL), pages 843–861.	826
774	Ilya Loshchilov and Frank Hutter. 2017. Decou-	Mariano Sigman and Guillermo A Cecchi. 2002. Global	827
775	pled weight decay regularization. <i>arXiv preprint</i>	organization of the wordnet lexicon. <i>Proceedings of</i>	828
776	<i>arXiv:1711.05101</i> .	<i>the National Academy of Sciences</i> , 99(3):1742–1747.	829
777	Andrea Madotto, Zhaojiang Lin, Genta Indra Winata,	Irene Solaiman, Miles Brundage, Jack Clark, Amanda	830
778	and Pascale Fung. 2021. Few-shot bot: Prompt-	Askeff, Ariel Herbert-Voss, Jeff Wu, Alec Rad-	831
779	based learning for dialogue systems. <i>ArXiv</i> ,	ford, Gretchen Krueger, Jong Wook Kim, Sarah	832
780	abs/2110.08118.	Kreps, et al. 2019. Release strategies and the so-	833
		cial impacts of language models. <i>arXiv preprint</i>	834
781	Fragkiskos D Malliaros and Konstantinos Skianis. 2015.	<i>arXiv:1908.09203</i> .	835
782	Graph-based term weighting for text categorization.		
783	In <i>Proceedings of the 2015 IEEE/ACM international</i>	Xiaohui Song, Longtao Huang, Hui Xue, and Songlin	836
784	<i>conference on advances in social networks analysis</i>	Hu. 2022. Supervised prototypical contrastive learn-	837
785	<i>and mining 2015</i> , pages 1473–1479.	ing for emotion recognition in conversation. <i>arXiv</i>	838
		<i>preprint arXiv:2210.08713</i> .	839
786	William C Mann and Sandra A Thompson. 1987.	Mark Steyvers and Joshua B Tenenbaum. 2005. The	840
787	<i>Rhetorical structure theory: A theory of text organiza-</i>	large-scale structure of semantic networks: Statistical	841
788	<i>tion</i> . University of Southern California, Information	analyses and a model of semantic growth. <i>Cognitive</i>	842
789	Sciences Institute Los Angeles.	<i>science</i> , 29(1):41–78.	843
790	Chris Mellish, Alistair Knott, Jon Oberlander, and	Yixuan Su, Fangyu Liu, Zaiqiao Meng, Tian Lan, Lei	844
791	Mick O’Donnell. 1998. Experiments using stochas-	Shu, Ehsan Shareghi, and Nigel Collier. 2021. Tactl:	845
792	tic search for text planning. In <i>Proceedings of the</i>	Improving bert pre-training with token-aware con-	846
793	<i>9th International General Workshop</i> , pages 98–107.	trastive learning. <i>arXiv preprint arXiv:2111.04198</i> .	847
794	ACL Anthology.		
795	Rada Mihalcea and Paul Tarau. 2004. Texttrank: Bring-	Enhua Tan, Lei Guo, Songqing Chen, Xiaodong Zhang,	848
796	ing order into text. In <i>Proceedings of the 2004 con-</i>	and Yihong Zhao. 2012. Spammer behavior analysis	849
797	<i>ference on empirical methods in natural language</i>	and detection in user generated content on social net-	850
798	<i>processing</i> , pages 404–411.	works. In <i>2012 IEEE 32nd International Conference</i>	851
		<i>on Distributed Computing Systems</i> , pages 305–314.	852
799	Marvin Minsky. 1982. Semantic information process-	IEEE.	853
800	ing.		
801	Adilson E Motter, Alessandro PS De Moura, Ying-	Peter D Turney. 2002. Learning to extract keyphrases	854
802	Cheng Lai, and Partha Dasgupta. 2002. Topology of	from text. <i>arXiv preprint cs/0212013</i> .	855
803	the conceptual network of language. <i>Physical Review</i>	Adaku Uchendu, Jeffrey Cao, Qiaozhi Wang, Bo Luo,	856
804	<i>E</i> , 65(6):065102.	and Dongwon Lee. 2019. Characterizing man-made	857
		vs. machine-made chatbot dialogs. In <i>TTO</i> .	858
805	Long Ouyang, Jeff Wu, Xu Jiang, Diogo Almeida, Car-	Adaku Uchendu, Zeyu Ma, Thai Le, Rui Zhang, and	859
806	roll L Wainwright, Pamela Mishkin, Chong Zhang,	Dongwon Lee. 2021. TURINGBENCH: A bench-	860
807	Sandhini Agarwal, Katarina Slama, Alex Ray, et al.	mark environment for Turing test in the age of neu-	861
808	2022. Training language models to follow in-	ral text generation . In <i>Findings of the Association</i>	862
809	structions with human feedback. <i>arXiv preprint</i>	<i>for Computational Linguistics: EMNLP 2021</i> , pages	863
810	<i>arXiv:2203.02155</i> .	2001–2016, Punta Cana, Dominican Republic. Asso-	864
		ciation for Computational Linguistics.	865
811	Matthew E Peters, Waleed Ammar, Chandra Bhaga-	Feng Wang and Huaping Liu. 2021. Understanding	866
812	vatula, and Russell Power. 2017. Semi-supervised	the behaviour of contrastive loss. In <i>Proceedings of</i>	867
813	sequence tagging with bidirectional language models.	<i>the IEEE/CVF conference on computer vision and</i>	868
814	In <i>Proceedings of the 55th Annual Meeting of the</i>	<i>pattern recognition</i> , pages 2495–2504.	869
815	<i>Association for Computational Linguistics (Volume</i>		
816	<i>1: Long Papers)</i> , pages 1756–1765.	Max Welling and Thomas N Kipf. 2016. Semi-	870
817	Alec Radford, Jeffrey Wu, Rewon Child, David Luan,	supervised classification with graph convolutional	871
818	Dario Amodei, Ilya Sutskever, et al. 2019. Language	networks. In <i>J. International Conference on Learn-</i>	872
819	models are unsupervised multitask learners. <i>OpenAI</i>	<i>ing Representations (ICLR 2017)</i> .	873
820	<i>blog</i> , 1(8):9.		

Dominic Widdows and Beate Dorow. 2002. A graph model for unsupervised lexical acquisition. In *COLING 2002: The 19th International Conference on Computational Linguistics*.

Yuta Yanagi, Ryohei Orihara, Yuichi Sei, Yasuyuki Tahara, and Akihiko Ohsuga. 2020. Fake news detection with generated comments for news articles. In *2020 IEEE 24th International Conference on Intelligent Engineering Systems (INES)*, pages 85–90. IEEE.

Zhilin Yang, Zihang Dai, Yiming Yang, Jaime Carbonell, Russ R Salakhutdinov, and Quoc V Le. 2019. Xlnet: Generalized autoregressive pretraining for language understanding. *Advances in neural information processing systems*, 32.

Liang Yao, Chengsheng Mao, and Yuan Luo. 2019. Graph convolutional networks for text classification. In *Proceedings of the AAAI conference on artificial intelligence*, volume 33, pages 7370–7377.

Rowan Zellers, Ari Holtzman, Hannah Rashkin, Yonatan Bisk, Ali Farhadi, Franziska Roesner, and Yejin Choi. 2019. Defending against neural fake news. *Advances in neural information processing systems*, 32.

Wanjun Zhong, Duyu Tang, Zenan Xu, Ruize Wang, Nan Duan, Ming Zhou, Jiahai Wang, and Jian Yin. 2020. Neural deepfake detection with factual structure of text. In *Proceedings of the 2020 Conference on Empirical Methods in Natural Language Processing (EMNLP)*, pages 2461–2470.

A Appendix

A.1 Details of GPT-3 Dataset

GPT-3 Dataset for CoCo is our latest dataset for the MGT detection task. There are two subsets in the self-made dataset for easy analysis of the impact of provenance and writing styles: unmixed- and mixed provinces. We use the text-davinci-003 model of OpenAI to generate MGT examples. The maximum length of HWTs is 1024 tokens, and the target generation length is set as 1024 tokens. Here is an example of the MGT data.

```
"title": "On Eve of World Cup, FIFA Chief Says, 'Don't Criticize Qatar; Criticize Me.'",
"text": "DOHA, Qatar. The president of world soccer's governing body on Saturday sought to blunt mounting concerns about the World Cup in Qatar with a strident defense of both the host country's reputation and FIFA's authority over its showpiece championship..... Citing statistics, history and even childhood to bolster his case, he at one point likened his own experience as a redheaded child of immigrants to Switzerland to the assimilation problems of gays in the Middle East, and defended the laws, customs and honor of the host country.",
"authors": ["Tariq Panja"],
"publish_date": "2022-11-19 00:00:00",
"source": "The New York Times",
"url": "https://www.nytimes.com/2022/11/19/sports/soccer/world-cup-gianni-infantino-fifa.html"
```

And the following data shows the corresponding MGT in the dataset.

```
"title": "On Eve of World Cup, FIFA Chief Says, 'Don't Criticize Qatar; Criticize Me.'",
"text": "The 2022 FIFA World Cup in Qatar is fast approaching, and its organizing committee's president, Gianni Infantino, is speaking out about the lingering criticism of the country hosting the event. .... he said. 'It is a once-in-a-lifetime opportunity for the region to show the world its values and aspirations, and it is vital that this event is seen as a celebration of football and a celebration of the region.'",
"authors": "machine",
"source": "The New York Times",
"matched_hwt_id": 202,
"label": "machine"
```

A.1.1 Human Written Texts

Unmixed Subset. The HWTs of the unmixed subset are all from The New York Times⁶ to exclude the impact of writing style. The time span of our data is Nov 1, 2022 - Dec 25, 2022, making sure that no pre-trained model has learned them. We develop the crawler based on news-crawler⁷.

Mixed Subset. The HWTs of the mixed subset come from various sources, listed as Table 5. The time span of the data is Jan 1, 2022 - Jan 7, 2023. We develop the crawler based on Newspaper3k⁸.

The dataset is specifically designed for MGTs detection and improving generation models. The contents of dataset are obtained from official news websites and the names of individual people are not mentioned maliciously. And we strongly reject using our dataset to create offensive content or peek at private information.

⁶<https://www.nytimes.com/>

⁷<https://github.com/LuChang-CS/news-crawler>

⁸<https://github.com/codelucas/newspaper>

Name	Website
Kotaku	https://kotaku.com
The Daily World	https://www.thedailyworld.com
CNN	https://edition.cnn.com
BBC	https://www.bbc.com
NBC News	https://www.nbcnews.com
Reuters	https://www.reuters.com
Huffpost	https://www.huffpost.com
Pando	http://pandodaily.com
Yahoo	https://news.yahoo.com
Sun Times	https://chicago.suntimes.com/news
Sfgate	https://www.sfgate.com
New Republic	https://newrepublic.com
Time	https://time.com
Pcmag	http://www.pcmag.com
CNBC	https://www.cnbc.com/world/
News	https://www.news.com.au/
The Atlantic	https://www.theatlantic.com/latest/

Table 5: Data sources for the mixed subset.

A.1.2 Machine Generated Texts

As the GPT-3 and ChatGPT model need prompts to generate, we write hints for the generation models to generate texts that meet our news-style long text generation. The hints format is as follows, and the content is related to HWTs.

```
Write a news more than 1000 words.
The news is written by {Authors} from {Source}
in {date}. Title is {title}.
```

A.2 GPT-3 Dataset Generated by Different Prompts and Experiment Results

To further validate the conclusion that GPT-3 generated texts are easier to detect, we utilize CNN news as reference and design different prompts for GPT-3 generation. The principle is to provide as more information as possible to GPT-3 for alleviating the possible gap in semantics and in length.

Keywords as Prompt (KP). We extract the keywords and entities with GPT-3.5-turbo and provide examples in original news to form the prompt for generation. The prompt format is as follows.

Example prompt for generation.

```
"role": "system", "content": "Extract all
the keywords, entities, and examples in the
following passage:"
"role": "user", "content": {text}
```

Example prompt for generation.

```
Generate a news passage.
The news is written by {Authors} from {Source}
in {date}.
Title: Lionel Messi isn't expected to be back
with PSG until early January after World Cup
success
Keywords: exploring, mountains, space, Poorna
Malavath, Kavya Manyapu, NASA, Mount Everest,
Project Shakthi, girls' education, Ladakh,
India, virgin peak, climbing, altitude sickness,
safety, motivation, empowerment, education,
gender gap, Mount Aconcagua, sponsorship.
Entities: CNN, Poorna Malavath, Kavya Manyapu,
NASA, Mount Everest, Project Shakthi, Ladakh,
India, Mount Aconcagua, South America, World
Bank.
Examples: designing space suits, youngest
ever woman to summit Mount Everest, climbed a
6,012m virgin peak, raise money to fund girls'
education, difficulties of climbing a virgin
peak, experiences of altitude sickness, purpose
of Project Shakthi, India's Right to Education
Act, sponsorship for underprivileged school
children, scaling Mount Aconcagua, expanding
sponsorship globally.
The target length for generation is 731 tokens.
Add as much details and examples as you can.
News:
```

Summary as Prompt (SP). We employ GPT-3.5-turbo to summarize the original texts. The compression ratio is set to [0.3, 1.0], which means the summary is required to be longer than 0.3 of the length of original text and shorter than whole original text. The generated summary is used as prompt and the format is as follows:

```
Generate a news based on the following
abstract:
Paris Saint-Germain's coach Christophe Galtier
has stated that Lionel Messi is not expected
to join the team until early January as he is
spending time in Argentina following the World
Cup. Kylian Mbappé, Neymar Jr. and Achraf
Hakimi, who played for their respective national
teams at Qatar 2022, could return to the team as
long as they are physically and mentally fit...
The news is written by Matias Grez from CNN in
2022-12-28 00:00:00.
Title: Lionel Messi isn't expected to be back
with PSG until early January after World Cup
success
News:
```

Outline as Prompt (OP). We also outline the skeleton of original texts by GPT-3.5-turbo and feed the outline into GPT-3 text-davinci-003. The prompt format is as follows:

Prompt for extraction.

```
"role": "system", "content": "Write a
hierarchical multi-point outline for the
paragraph."
"role": "user", "content": {text}
```

Example prompt for generation.

News Title: There’s a shortage of truckers, but TuSimple thinks it has a solution: no driver needed
The news is written by Jacopo Prisco, CNN from CNN in 2021-07-15 02:46:59.
Outline:
I. TuSimple’s plan for fully autonomous truck tests
A. Reliability of software and hardware needs to improve
B. Fully autonomous tests without human safety driver planned by end of year
C. Results will determine if company can launch trucks by 2024
D. 7,000 trucks reserved in US alone
II. TuSimple’s competition
A. ...
Add more details and examples.
News:

We first remove the HWTs that do not have desired length (i.e., 200-1024 tokens). And we take half of the selected HWTs as references to formulate different prompts mentioned above and feed it into GPT-3 to get MGTs. The MGTs are sampled by Gaussian Distribution of their lengths. To avoid the possible label leakage brought by text length, we directly filter the no-reference HWTs according to the Gaussian Distribution of MGT lengths.

Besides the self-constructed datasets, we also utilize the published GPT-3 dataset TuringBench benchmark (abbreviate as GPT-3 (TB)) (Uchendu et al., 2021) to validate the deceptiveness of GPT-3. The statistics of datasets we use is in Table 6.

Dataset		Train	Valid	Test	# of tokens
GPT-3(KP)	HWT	446	148	148	427.96 $\pm$ 45.49
	MGT	446	148	148	403.88 $\pm$ 75.63
GPT-3(SP)	HWT	446	148	148	427.96 $\pm$ 45.49
	MGT	446	148	148	415.72 $\pm$ 66.54
GPT-3(OP)	HWT	446	148	148	427.96 $\pm$ 45.49
	MGT	446	148	148	429.34 $\pm$ 78.62
GPT-3(TB)	HWT	5,964	975	1915	236.17 $\pm$ 72.96
	MGT	5,507	894	1763	147.29 $\pm$ 70.15

Table 6: Statistics of GPT-3 datasets.

We conduct experiments with 3 random seeds and the average results are shown in Table 7. Counterintuitively, even if we elaborate the prompts and eliminate the length difference between MGTs and HWTs, the detection results are still superior, even on outdated baselines like GPT-2. The conclusion might be counterintuitive, but texts generated by the most advanced and popular GPT-3 model are the easiest to detect.

A.3 Implementation Details

This part mentions the implementation details and hyper-parameter settings of all the methods in the

experiment. To imitate the situation of low data-resources, we sample 10% texts from the datasets as limited dataset, which will test models together with the complete datasets. And we conduct experiment on 10 different seeds and report the average test accuracy, F1-Score, and standard deviation.

We use RoBERTa base model to initialize the embedding of our representation and optimize the model using AdamW (Loshchilov and Hutter, 2017) optimizer with a 0.01 weight decay. We set the initial learning rate to 10^{-5} and the batch size to 8 for all datasets based on experiences.

We utilize packages, namely transformers, pytorch, and allennlp to implement CoCo. And the GPT-3 datasets and ChatGPT case is generated by OpenAI API and websites. We spend \$300 for API costs, including development and final generation costs. We train and do experiments on 8 NVIDIA A100 GPUs on 2 Ubuntu-based servers. The total budget for training 20 epochs, dev, and testing on the GROVER dataset is 2.5 hours. On GPT-2 dataset is 12 hours, and on GPT-3 dataset is 1.5 hours. We will publish our code and dataset recently.

A.4 Effect of Hyper-Parameters

A.4.1 Contrastive Learning Parameters

We evaluate the influence of contrastive learning hyper-parameters α and τ with experiments on different combinations of them. The result is shown in Fig. 5. Considering the discovering that smaller τ leads to better hard negative mining ability (Wang and Liu, 2021), we select α from $\{0.1, 0.2, \dots, 0.9\}$ and τ from $\{0.1, 0.2, 0.3\}$. We find that the extreme α value causes the performance degradation and the best hyper-parameter combination is $\alpha, \tau = 0.6, 0.2$. Our analysis is that large α forces the model to concentrate on the instance-level contrast and small α lets class separation objective take control. Both will reduce the generalization performance of the detector on test set.

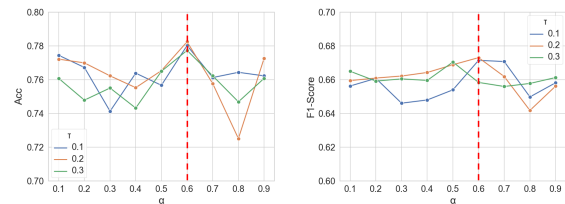


Figure 5: Effect of parameters α and τ on model performance.

Dataset	GPT-3 (KP)		GPT-3 (SP)		GPT-3 (OP)		GPT-3 (TB)	
Metric	ACC(val/test)	F1(val/test)	ACC (val/test)	F1 (val/test)	ACC (val/test)	F1 (val/test)	ACC (val/test)	F1 (val/test)
GPT2	0.9914/0.9916	0.9916/0.9918	0.9890/0.9893	0.9885/0.9889	0.9925/0.9928	0.9923/0.9924	0.9884/0.5422*	0.9880/0.6335*
RoBERTa	0.9946/0.9950	0.9950/0.9952	0.9935/0.9941	0.9933/0.9937	0.9946/0.9943	0.9942/0.9940	0.9962/0.6406*	0.9960/0.7273*
CoCo	0.9955/0.9950	0.9942/0.9945	0.9938/0.9941	0.9936/0.9940	0.9942/0.9943	0.9942/0.9943	0.9966*	0.9970*

Table 7: Experiment of different detectors on different GPT-3 Dataset. *:The great performance difference between validation set and test set on GPT-3 (TB) are because the test set randomly sample 50% of the words of each article in the dataset (Uchendu et al., 2021). We do not test CoCo on GPT-3 (TB) for the reason that such operation greatly influences the coherence in texts. We provide an example of this in Table 8.

GPT-3 (TB)	GPT-3 (OP)
'video : morne morkel press conference * cricbuzz.video : england cricbuzz.bevan leads scotland 's 21-man squad for their first ever test match against pakistan in edinburgh icc.chris rogers retires after champions trophy defeat : australian cricketer announces international retirement the sun.icc super eight teams : odi ranking results.bahrain host oman on sunday kitply hans vohra gold cup gulf today.icc results.new zealand series history : india v new zealandyazan mohsen qawasma : how bahrain caught	Recent changes to key international indexes have resulted in the unprecedented exclusion of Russian stocks at a "zero" price, causing further losses in Moscow's already-dismal stock exchange. This exclusion has made Russia no longer an option for investors, prompting a shift to other emerging markets.\n\nThe dramatic shift was made in early March, when FTSE Russell and MSCI announced the removal of Russian stocks from their indexes due to the country's escalating economic and geopolitical problems. Shortly after, the Moscow Exchange suspended trading, sending ripples through the market.\n\nThe possible default on Russian debt has Western investors further reconsidering their investments in Russia...

Table 8: A comparison example between texts in test set of GPT-3 (TB) and GPT-3 (OP). The GPT-3 (TB) text shows great disorder while GPT-3 (OP) text is neat.

A.4.2 Graph Parameters

We further investigate the effect of max node number and max sentence number on model performance. The result is shown in Fig. 6. We select max node number from {60, 90, 120, 150} and max sentence number from {30, 45, 60, 75}. The detector performs best when max node number is 90 and max sentence number is 45. The experiment results prove that the large node and sentence number are not necessary for the improvement of detection accuracy. We infer that even though setting large node and sentence number includes more entity information, excessive nodes bring noise to the model and impair the distinguishability of coherence feature.

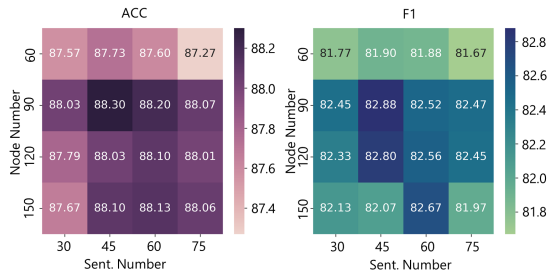


Figure 6: Performance of CoCo with different graph parameters.

A.5 Static Geometric Analysis on Coherence Graph

We have witnessed performance enhancement by applying the graph-based coherence model to the detection model, but how does the coherence graph help detection? In this subsection, we apply static geometric features analysis to coherence graph we construct to evaluate the distinguishable difference between HWTs and MGTs with explanation. In the following discussion, we take the dataset of GROVER into the analysis. Some basic metrics of data and the corresponding graph are shown in Table 9.

Metric	HWT	MGT
Sample Num.	4994	4991
Avg. Num. of Token	463.2	456.0
Avg. Num. of Vertex	43.60	32.37
Avg. Num. of Edge	107.4	65.44

Table 9: Basic metrics of texts and corresponding graphs.

Though HWTs and MGTs have approximately the same number of tokens in every text, coherence graph for HWTs has larger scale than MGTs'

Metric	Avg. Degree
HWT	2.980
MGT	2.591

Table 10: Average of degree (whole dataset).

with **34.7%** more vertexes and **64.1%** more edges, which shows that HWTs have more complex semantic relation structures than MGTs.

A.5.1 Degree Distribution

Semantically, degree of coherence graph measures the co-occurrence and TF-IDF feature of keywords. Moreover, degree distribution shows global coherence because high-degree nodes devote to the main topic and low-degree nodes are the extension.

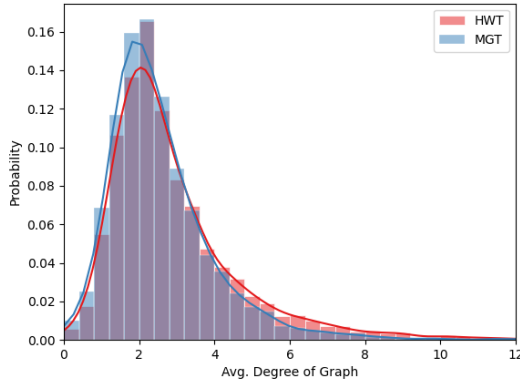


Figure 7: Distribution of average degree of graphs.

As shown in Table 10, the degree of the graph representation of HWTs is **15.0%** larger than MGTs, which shows disparities of MGTs to form coherent interaction between sentences. Fig. 7 measures the distribution of each graph’s average nodes’ degree, showing that the distribution of HWTs has a longer tail than MGTs.

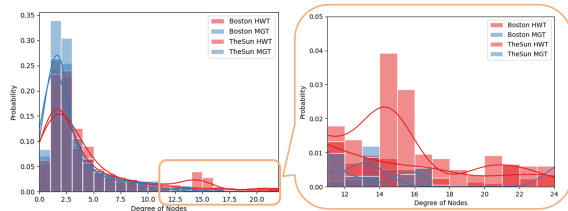


Figure 8: Distribution of degree with different provenance.

Furthermore, we analyze the distinguishability of degree features when impacted by other factors.

One most considerable influences is the style and genre of different provenance. We chose around 60 articles from The Sun⁹ and Boston¹⁰. Then we use GROVER to mimic their style to generate similar topic news. Fig. 8 shows the degree distribution of HWTs and MGTs of both provenances.

We use Jensen–Shannon divergence to evaluate the similarity of the degree distribution. The JS-divergence of MGTs mimicking The Sun and Boston is **0.029**, while the JS-divergence of MGTs and HWTs in Boston is **0.050**, in The Sun is **0.061**. The apparent gap shows that degree distribution can robustly detect MGTs and HWTs when impacted by provenance differences.

A.5.2 Aggregation

Aggregation is a shared metric for complex networks and linguistics, depicting how closely the whole is organized around its core. We propose two metrics to evaluate the aggregation of graph-based text representation in our coherence model, the size of the largest connected subgraph and the clustering coefficient.

In our representation, not all sentences have entities related to others. Hence the graph is an unconnected one. The average number of nodes in subgraphs of MGTs is **4.49** and of HWTs is **4.84**. We propose that the size of the largest connected subgraph shows the contents which are closely organized around the topic. Moreover, the size of graphs may be an unfair factor, so we use the portion of nodes in the largest connected subgraph to reflect its size. The average portion in HWTs is **0.6725** and in MGTs is **0.6458**. Fig. 9 shows the distribution of the portion of graphs, and HWTs distribute more high-portion ones than MGTs.

The clustering coefficient represents how nodes tend to cluster. For the entities of texts, clustering evaluates how the author narrates around the central theme. The larger the clustering coefficient is, the tighter the semantic structure is. The average cluster coefficient of the graphs of HWTs is **0.2213** and of MGTs is **0.1983**, HWTs is **11.6%** better than MGTs. Fig. 10 shows the distribution.

A.5.3 Core & Degeneracy

The degeneracy of a graph is a measure of how sparse it is, and the k -core is the subgraph corresponding to its significance in the graph. We

⁹<https://www.thesun.co.uk/>

¹⁰<https://www.boston.com/>

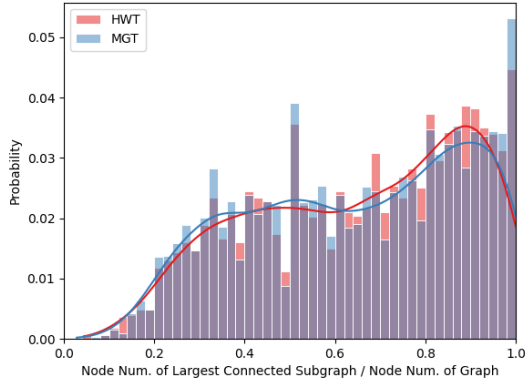


Figure 9: Portion of the largest connected subgraph.

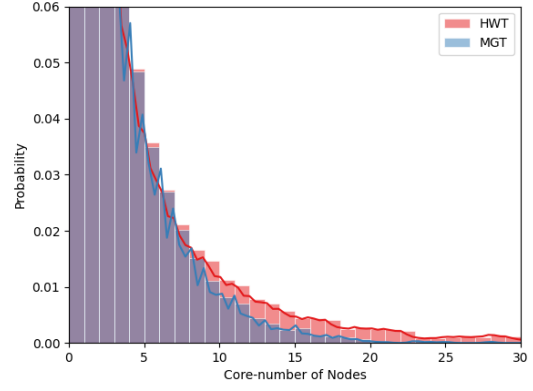


Figure 11: Core-number of nodes in graphs

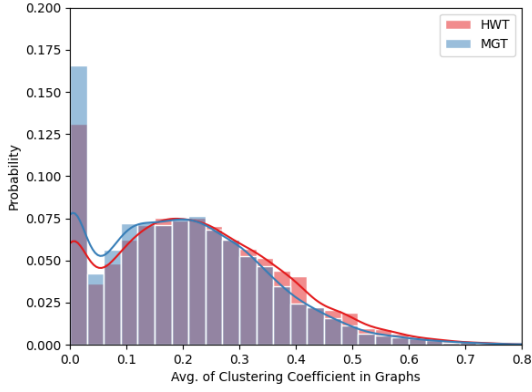


Figure 10: Distribution of clustering coefficient.

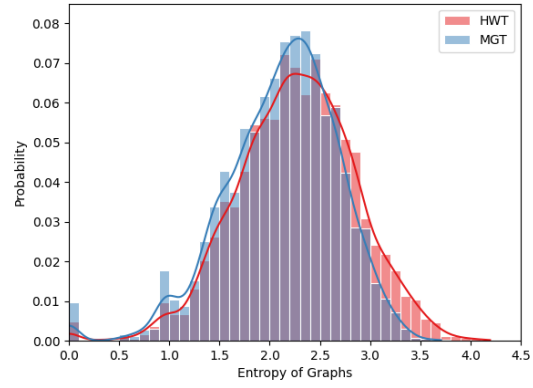


Figure 12: Structure entropy of graphs

propose that, in our graph representation, the degeneracy process of graphs equals summarizing texts semantically. The maximum of core-number shows the complexity of hierarchical structure in texts. Furthermore, the distribution of the core-number reflects the overall sparse and is a graph-perspective N-gram module. Based on experiments, the average core-number of HWTs is **5.772** while MGTs with **4.458**. HWTs are **29.5%** ahead. Fig. 11 is the distribution of the core-number.

A.5.4 Entropy

Entropy is a scientific concept to measure a state of disorder, randomness, or uncertainty. The well-known Shannon entropy is the core of the information theory, measuring the self-information content. For the graph data, network structure entropy defined as the following can examine the information amount of the graph structure.

$$Entropy = - \sum_{i=1}^N I_i \ln I_i = - \sum_{i=1}^N \frac{k_i}{\sum_{j=1}^N k_j} \ln \left(\frac{k_i}{\sum_{j=1}^N k_j} \right), \quad (9)$$

where I_i is the information content represented by the degree distribution, N is the number of nodes, and k_i is the degree of the i -th node.

Global coherence, from our perspective, equals refining more information inside the semantic structure of the whole text, which matches to structure entropy of our graph representation. From our experiments, the structure entropy of HWTs (2.263) is **6.80%** larger than MGTs (2.119), which means HWTs obtain more structured information because their semantic information is globally organized. We show the network structure entropy distribution in Fig. 12.

A.6 Exploration on Imbalanced Data

Imbalanced distribution in data is another crucial limitation in the task of MGTs detection, which is similar to the low resource limitation. It is imaginable that, with the development of generation technology, MGTs will overwhelmingly dominate low-quality articles since they are easier and faster to generate than human writing. The detection model will face training resources with MGTs as

the main part and HWTs as the small part. We test the current models in the imbalanced limitation and find the dramatic decline in accuracy when the ratio of HWTs is less than 30%, as shown in the Fig. 13. The test is based on the 10% GROVER dataset.

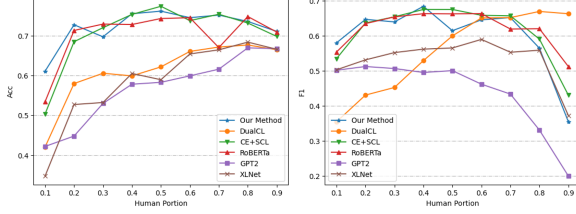


Figure 13: Model comparison results on DL dataset with 9 different human-generated text portions.

All models show poor performance at low HWTs ratios. With a percentage of HWTs of 0.1 (only 100 HWTs in the training set in this case), most of the models have an accuracy below 50%, which performance is close to random and reflects intolerance for extreme cases. Besides, we find that a high proportion of HWTs also cause a decrease in F1 score to some extent.

A.7 Related Work: Graph-based Text Representation

Graph can represent text structure and inner-relation (Minsky, 1982). Based on different organizing methods, graphs can reflect static statistical (e.g., co-occurrence (Cancho and Solé, 2001), collocation (Bordag et al., 2003)), dynamically statistical (e.g., evolution (Dorogovtsev and Mendes, 2001)), lexical (Widdows and Dorow, 2002), orthographic (Choudhury et al., 2007), cognitive (e.g., conception (Motter et al., 2002)), syntactic (Cancho et al., 2004; Ferrer Cancho et al., 2007), semantic (Steyvers and Tenenbaum, 2005; Sigman and Cecchi, 2002; Kozareva et al., 2008) relations. Hasan et al. (2017) propose an overall survey about graph-based text representation.

Graph-of Words (GoW) Model (Turney, 2002; Mihalcea and Tarau, 2004) is a type graph representation method in which each document is represented by a graph, whose nodes correspond to terms and edges capture co-occurrence relationships between terms. Using GoW, keywords can be extracted by retaining the document graph (Turney, 2002). Thus, graph representation is sensible to apply in tasks like information retrieval (Blanco and Lioma, 2011), categorization (Malliaros and Skianis, 2015) and sentiment classification tasks

(Huang and Carley, 2019; Hou et al., 2021). Most models enhance classification or detection performance by combining graph representation with neural networks. Text-GCN (Yao et al., 2019) first builds a single large graph for the whole corpus, followed by Tensor-GCN (Liu et al., 2020) with tensor representation. Also, the relation between words varies, and should be treated as different edges. CoCo matches keywords PLM embedding to nodes and sentence representation, considers dealing inner- and inter-sentence relation differently in GCN, and merges the structure graph and flat sequence representation to predict accurately.

A.8 Pseudocode of CoCo

Algorithm 1 Algorithm of CoCo

Input: Input X , consisting of documents D and corresponding coherence graph G , hyperparameters such as the size of dynamic memory bank M and batch size S , labels Y

Output: A learned model CoCo, consisting of key encoder f_k with parameters θ_k , query encoder f_q with parameters θ_q , classifier f_c with parameters θ_c

- 1: Initialize $\theta_k = \theta_q, \theta_c$
- 2: Initialize dynamic memory bank with $f_k(x_1, x_2 \dots x_M)$, where x_i is randomly sampled from X .
- 3: Freeze θ_k
- 4: $epoch \leftarrow 0$
- 5: **while** $epoch \leq epoch_{max}$ **do**
- 6: $n \leftarrow 0$
- 7: **while** $n \leq n_{max}$ **do**
- 8: Randomly select batch b_k, b_q
- 9: $D_q = f_q(b_q), D_k = f_k(b_k)$
- 10: $\hat{p} = f_c(D_q)$
- 11: Calculate $\mathcal{L}_{ICL}$ with equation 5, calculate $\mathcal{L}_{CE}$ with equation 6, calculate $\mathcal{L}_{total}$ with equation 7
- 12: Backward on $\mathcal{L}_{total}$ and update θ_q, θ_c based on AdamW gradient descent with an adjustable learning rate
- 13: Momentum update θ_k with equation 8
- 14: Update dynamic memory bank $queue$ with $enqueue(queue, D_k), dequeue(queue)$
- 15: $k \leftarrow k + 1$
- 16: **end while**
- 17: **if** Early stopping **then**
- 18: **break**
- 19: **else**
- 20: $epoch \leftarrow epoch + 1$
- 21: **end if**
- 22: **end while**
- 23: **return** A trained model CoCo
